



Leading the Human-AI Paradox: Paradoxical Leadership, Human-AI Complementarity, Responsible AI Agency, and Employee Innovation

Samreen Batool¹ , Ahmad Adeel²

Article History:

Received: 22-05-2026
Revision: 05-08-2026
Accepted: 07-09-2026
Publication: 08-09-2026

Cite this article as:

Batool, S., Adeel, A. (2026). Leading the Human-AI Paradox: Paradoxical Leadership, Human-AI Complementarity, Responsible AI Agency, and Employee Innovation. *Innovation Management Frontiers*, 1(1), 47-55. doi.org/10.36923/imf.v1i1.432

©2026 by author(s). This is an open-access article distributed under the terms of the Creative Commons Attribution License 4.0 International.

Corresponding Author:

Adeel Ahmad

Faculty of Business and Communications, INTI International University, Nilai, Malaysia. Email: adeelleads@yahoo.com

Abstract: The increasing use of generative artificial intelligence (GenAI) in organizational work creates a persistent tension between leveraging machine capabilities and preserving human judgment, autonomy, creativity, and accountability. Drawing on paradox theory, this conceptual paper develops a framework explaining how paradoxical leadership can enable organizations to manage this human-AI tension through a both-and approach. The framework proposes two complementary mechanisms. Human-AI complementarity captures the productive integration of distinct human and AI capabilities, whereas responsible AI agency reflects employees' capacity to use AI while retaining critical evaluation, meaningful decision authority, and accountability. The framework also identifies AI literacy as an enabling boundary condition and AI overreliance as a constraining condition that influences whether human-AI collaboration results in augmentation or substitution. By integrating paradoxical leadership, human-AI collaboration, responsible AI, and innovation research, the paper extends paradox theory to AI-enabled work and explains why GenAI may enhance employee innovation under some conditions while generating dependence and diminished human contribution under others. The paper develops 7 propositions to guide future empirical research on leadership in increasingly hybrid human-AI workplaces.

Keywords: Paradoxical Leadership, Generative Artificial Intelligence, Human-AI Complementarity, Responsible AI Agency, AI Literacy, AI Overreliance, Innovative Work Behavior

1. Introduction

Generative artificial intelligence (GenAI) is rapidly altering the organization of knowledge work by extending employees' capacity to generate ideas, process information, solve problems, and make decisions. However, the growing integration of intelligent systems into organizational work is not a straightforward technological substitution for human effort. Rather, it creates an enduring tension between automation and augmentation because the same technologies that enhance productivity and expand human capabilities can also displace human input, encourage cognitive dependency, and alter employees' autonomy and professional roles (Raisch & Krakowski, 2021; Ritala, Ruokonen, & Ramaul, 2024; Hillebrand, Raisch, & Schad, 2025; Fügner, Walzner, & Gupta, 2025; Adeel, Mohy-ul-din, Ashraf, & Sheikh, 2026). Consequently, contemporary organizations face a fundamental managerial challenge: they must leverage AI's efficiency, analytical power, and generative capabilities while preserving human judgment, creativity, autonomy, and responsibility.

This challenge is particularly evident in innovation-related work. AI can augment employee creativity by expanding access to information, accelerating idea generation, and supporting cognitive processes that contribute to innovative outcomes (Jia, Luo, Fang, & Liao, 2024). GenAI can also support different phases of creative work and improve innovative job performance through cognitive and social use, knowledge transfer, and resource acquisition (Huang, Long, Zhu, & Zhu, 2026; Zhang, Zhu, Zhang, & Shohruh, 2026). Yet these benefits are not automatic. AI dependence can suppress creativity, output homogenization can reduce collective diversity, and extensive reliance on machine-generated ideas can weaken originality or psychological ownership (Doshi & Hauser, 2024; Yin, Jiang, & Niu, 2024; Ma, Ge, Li, & Wang, 2026; Cui, Wang, Cao, & Zhu, 2026). GenAI therefore functions less as an unconditional source of innovation than as an innovation amplifier whose effects depend on how employees interact with, evaluate, and integrate AI-generated outputs.

The emerging literature on human-AI collaboration offers an important response by shifting attention from replacement toward complementarity. Human-AI complementarity emphasizes configuring distinctive human and AI capabilities so joint performance exceeds what either could achieve independently. AI contributes rapid information processing, pattern recognition, and option generation, whereas humans remain important for contextual interpretation, tacit knowledge, ethical reasoning, and judgment under ambiguity (Ren, Deng, & Joshi, 2023; Dittmar & Sposato, 2026; Adeel, Hussain, & Sheikh, 2026). Accordingly, effective human-AI work depends on task allocation, calibrated trust, role clarity, and coordination rather than indiscriminate delegation to AI (Jain, Garg, & Khera, 2023; Kumar, Kumar, & Raj, 2025; Aman, Hamid, Mat-Isa, & Mohamad, 2025).

Nevertheless, complementarity alone does not resolve the managerial challenge created by GenAI. Employees may collaborate extensively with AI while gradually relinquishing independent judgment, decision authority, or responsibility. Generative AI can simultaneously

¹ School of Management, Universiti Sains Malaysia (USM), Minden, 11800 Penang, Malaysia

² Faculty of Business and Communications, INTI International University, Nilai, Malaysia

enable and constrain human agency: it can reduce cognitive burden and support higher-order work while also encouraging cognitive outsourcing, dependency, deskilling, and reduced critical reflection when outputs are accepted uncritically (Di Plinio, 2025; Elmratty, Bani-Hani, & Alawadi, 2026; Özer, Perc, & Özçelik, 2026). Excessive reliance on AI is also associated with automation bias and poorer discrimination between appropriate and inappropriate recommendations, particularly when trust is miscalibrated (Goddard, Roudsari, & Wyatt, 2012; Klingbeil, Grützner, & Schreck, 2024; Ahmad & Mehmood, 2025; Romeo & Conti, 2026). Responsible AI scholarship consequently emphasizes human oversight, accountability, transparency, and the preservation of meaningful decision control (Papagiannidis, Mikalef, & Conboy, 2025; Novelli, Taddeo, & Floridi, 2024; Sturm, Krause, & van Giffen, 2026). Thus, the central organizational problem is not simply whether employees should use AI, but how organizations can simultaneously promote productive human-AI integration and preserve meaningful human agency.

Such simultaneous and interdependent demands make paradox theory particularly relevant. Rather than assuming that competing demands must be resolved by selecting one alternative, paradox theory emphasizes persistent contradictions whose opposing elements remain interdependent and therefore require a both-and approach (Lewis, 2000; Smith & Lewis, 2011; Schad, Lewis, Raisch, & Smith, 2016). This logic has long informed paradoxical leadership. Zhang, Waldman, Han, and Li (2015) conceptualized paradoxical leadership as the simultaneous enactment of apparently competing leadership behaviors, such as maintaining control while granting autonomy and enforcing requirements while allowing flexibility. Subsequently, Zhang, Zhang, and Law (2022) demonstrated that paradoxical leadership can enhance individual and team innovation by facilitating ambidexterity, showing how leaders enable employees and teams to accommodate competing demands rather than simply choosing between them.

The characteristics of GenAI-enabled work create a compelling setting for extending this perspective. Employees are increasingly expected to use AI intensively while remaining capable of independent judgment; exploit algorithmic efficiency while maintaining autonomy; draw on machine-generated ideas while preserving originality; and benefit from automation while remaining accountable for decisions. Recent paradox-oriented research similarly identifies tensions surrounding autonomy and control, creativity and dependence, efficiency and judgment, and governance and experimentation in GenAI contexts (Alshalfan, 2026; Choi, Park, Paydas Turan, & Lee, 2026; Papadonikolaki, 2026; Esposito De Falco, Laviola, Mercuri, & Cucari, 2026).

Despite this theoretical fit, an important gap remains between paradoxical leadership research and the emerging literature on human-AI work. Direct evidence is beginning to appear. Zhang and Wu (2026), for example, show that paradoxical leadership shapes the consequences of different forms of human-AI interaction, while paradox-oriented thinking has been associated with employee engagement in AI collaboration (Marvi, Foroudi, & AmirDadbar, 2025). However, the available evidence remains fragmented and does not yet explain how paradoxical leadership can transform human-AI tensions into productive complementarity while simultaneously preventing the erosion of human agency. At the same time, human-AI complementarity research and responsible AI research have largely developed as parallel conversations: one emphasizes capability integration and performance, whereas the other emphasizes oversight, accountability, and human control.

To address this gap, the present paper develops a paradox-based framework of leadership in human-GenAI collaboration. It proposes that paradoxical leadership enables employee innovative work behavior through two interrelated mechanisms: human-AI complementarity, through which employees combine distinct human and AI capabilities, and responsible AI agency, through which employees retain critical judgment, decision authority, and accountability while using AI. Responsible AI agency is introduced here as a conceptual construct rather than treated as an established scale. The framework further identifies AI literacy and AI overreliance as boundary conditions that respectively facilitate or constrain these processes. In doing so, the paper extends paradoxical leadership from predominantly human-centered organizational contradictions to the emerging tension between human agency and machine intelligence and provides a theoretical explanation for why GenAI may enhance innovation in some organizational settings while producing dependence and diminished human contribution in others.

2. Theoretical Background

2.1. Paradox Theory and Paradoxical Leadership

Organizations frequently confront demands that are individually legitimate but mutually contradictory. Paradox theory conceptualizes such conditions as the simultaneous presence of contradictory yet interrelated elements that persist over time rather than disappearing once managers select one alternative (Lewis, 2000; Smith & Lewis, 2011). Effective responses therefore require actors to acknowledge the coexistence of opposing demands and develop dynamic approaches that accommodate both sides. This both-and orientation is particularly relevant under complexity and uncertainty, where competing demands tend to remain salient rather than being permanently resolved (Schad et al., 2016).

Paradoxical leadership translates this logic into observable leader behavior. Zhang et al. (2015) identify five behavioral combinations: enforcing work requirements while allowing flexibility; maintaining decision control while permitting autonomy; treating employees equitably while recognizing individuality; retaining formal authority while sharing influence; and preserving hierarchical distance while maintaining interpersonal closeness. These combinations do not imply random switching between styles. Rather, paradoxical leadership reflects the capacity to integrate apparently contradictory but interdependent behaviors according to situational demands. Subsequent research shows that such leadership can strengthen paradox mindset, engagement, creativity, knowledge-related behavior, and innovation (Liu, Xu, & Zhang, 2020; Boemelburg, Zimmermann, & Palmié, 2023; Batool, Raziq, & Sarwar, 2023; Wei, Zhou, & Wang, 2023; Nagarajan, 2024).

Innovation is itself paradoxical. Employees require freedom to generate unconventional ideas while operating within organizational constraints; innovation also involves both idea generation and implementation and therefore benefits from divergent and convergent activities (Scott & Bruce, 1994; Zhang et al., 2022; Alamanda, Ahmed, Kurniady, Rahayu, Ahmad, & Hashim, 2024). Zhang et al. (2022) demonstrate that paradoxical leadership enhances individual and team innovation through individual and team ambidexterity. Their work shows that leaders can promote exploration through autonomy and flexibility while supporting exploitation through standards, coordination, and structural control. This establishes paradoxical leadership as an important mechanism for translating competing demands into innovative outcomes.

2.2. Generative AI and the Emerging Human-AI Paradox

AI-enabled work has long been characterized by a tension between automation and augmentation. Automation transfers activities from humans to machines, whereas augmentation emphasizes arrangements through which AI extends rather than replaces human capability (Raisch & Krakowski, 2021; Ahmad, Thurasamy, Adeel, & Alam, 2023). GenAI intensifies this tension by extending AI participation from routine tasks into analysis, ideation, communication, creativity, and decision support. Recent research consequently identifies enduring tensions involving autonomy versus control, efficiency versus reflection, creativity versus homogenization, trust versus oversight, and augmentation versus deskilling (Choi et al., 2026; Papadonikolaki, 2026; Cristofaro & Bañón-Gómis, 2026).

These tensions are not merely technological. AI can reduce cognitive burden and support decision-making, but extensive delegation can also encourage cognitive offloading and weaken independent judgment (García-Barrios, 2025; Elmrayyan et al., 2026). Similarly, GenAI can expand the range of ideas available to employees while simultaneously reducing diversity or originality when users converge around machine-generated suggestions (Doshi & Hauser, 2024; Lee & Chung, 2024; Huang et al., 2026). Trust is also paradoxical: too little trust prevents employees from exploiting useful AI outputs, whereas excessive trust increases the risk of inappropriate deference (Okamura & Yamada, 2020; Cheng, Dai, Liu, & Peng, 2026).

Accordingly, this paper conceptualizes GenAI-enabled work as a human-AI paradox: organizations and employees must simultaneously increase the productive use of AI capabilities while retaining human judgment, autonomy, creativity, and accountability. Paradox theory is particularly suited to this problem because maximizing one side can undermine the other. Restricting AI use may protect autonomy but sacrifice technological benefits; extensive delegation may improve short-term efficiency while eroding human contribution. The challenge therefore requires ongoing management rather than a stable either-or solution.

2.3. Human-AI Complementarity

Human-AI complementarity refers to the productive integration of distinct human and AI capabilities. It is grounded in the recognition that humans and AI possess different strengths and that joint performance depends on how those differences are configured across tasks (Ren et al., 2023; Dittmar & Sposato, 2026). AI is particularly suited to rapid information processing, pattern recognition, and large-scale comparison, whereas humans contribute contextual interpretation, tacit knowledge, ethical reasoning, and judgment under ambiguity. Complementarity therefore requires deliberate allocation of responsibilities rather than simple co-presence of human and machine intelligence. Importantly, complementarity is conditional. Effective human-AI collaboration depends on task fit, role clarity, coordination, transparency, user expertise, and calibrated trust (Jain et al., 2023; Kumar et al., 2025; Aman et al., 2025; Vössing, Kuhl, Lind, & Satzger, 2022). Miscalibrated trust can undermine collaboration because employees may either reject useful AI recommendations or accept inappropriate ones. Similarly, GenAI collaboration can improve creative performance but may reduce diversity or innovation when employees rely excessively on machine-generated suggestions (Doshi & Hauser, 2024; Mascareño, Wörtler, Przeglasińska, & Ciechanowski, 2026). Thus, human-AI complementarity concerns the quality of capability integration rather than the frequency or intensity of AI use.

2.4. Responsible AI Agency

Complementarity does not by itself establish whether employees remain meaningful decision makers. Human agency becomes essential because AI can simultaneously support and constrain employees' capacity for independent action. Decision support can reduce routine cognitive load and expand the information available to employees, while cognitive outsourcing, opacity, and excessive reliance can weaken reflection, autonomy, and judgment (Di Plinio, 2025; Elmrayyan et al., 2026; Özer et al., 2026). Human-in-the-loop research therefore emphasizes the continuing need for humans to review, validate, contextualize, and override machine recommendations (Lazaros, Vrahatis, & Kotsiantis, 2026). Responsible AI scholarship reaches a similar conclusion from a governance perspective. Transparency, explainability, accountability, traceability, and meaningful human oversight are repeatedly identified as central principles of responsible organizational AI use (Papagiannidis et al., 2025; Novelli et al., 2024; Lee, Yoon, & Lee, 2026; Sturm et al., 2026). Yet oversight matters only when employees have sufficient authority and competence to intervene, not merely provide symbolic approval. Building on these streams, this paper introduces responsible AI agency as employees' capacity to use AI productively while retaining critical evaluation, meaningful decision authority, and accountability for AI-assisted actions. Responsible AI agency does not imply resistance to AI. It involves active and reflective use in which employees remain capable of accepting, modifying, questioning, or rejecting AI-generated outputs according to situational requirements. Human-AI complementarity therefore captures capability integration, whereas responsible AI agency captures the preservation of human judgment and responsibility within that integration.

2.5. AI Literacy and AI Overreliance

The proposed mechanisms are unlikely to operate uniformly across employees. AI literacy is increasingly conceptualized as a multidimensional competency involving technical understanding, practical use, critical evaluation, and ethical awareness (Chee, Ahn, & Lee, 2025; Benlian & Pinski, 2025; Liu, Zhang, & Wei, 2025; Markus, Carolus, & Wienrich, 2025; Annapureddy, Fornaroli, & Gatica-Pérez, 2025). Such competence should enable employees to recognize both the capabilities and limitations of GenAI and to translate leadership guidance into productive collaboration. AI literacy has also been linked to job performance, innovation capacity, and innovative performance (Liu et al., 2025; Steinhäuser & Heid, 2026; Zhou, Lü, & Li, 2026; Tan, Nawaz, Shu, & Ramzan, 2026). By contrast, AI overreliance represents excessive deference to AI outputs. Automation-bias research shows that users may accept automated recommendations despite contradictory evidence, make omission or commission errors, and reduce independent verification when trust becomes excessive (Parasuraman & Manzey, 2010; Goddard et al., 2012; Alon-Barkat & Busuioc, 2023). More recent work similarly shows that excessive AI reliance can weaken critical judgment, increase conformity, and suppress creativity, even though appropriate reliance can improve performance under favorable task conditions (Klingbeil et al., 2024; Romeo & Conti, 2026; Cui et al., 2026). These contrasting conditions help explain whether human-AI collaboration remains augmentative or shifts toward substitution.

3. Conceptual Framework and Proposition Development

Building on paradox theory, the proposed framework explains how paradoxical leadership can help employees manage the tension between human agency and machine intelligence. The central argument is that GenAI-enabled work creates persistent and interdependent demands rather than a technological choice between human and artificial intelligence. Human-AI complementarity is the capability-integration mechanism through which distinct human and AI strengths combine, whereas responsible AI agency is the human-control mechanism through which employees retain judgment and accountability. The framework further proposes that these mechanisms jointly influence innovative work behavior and that their effects depend on AI literacy and AI overreliance.

3.1. Paradoxical Leadership and Human-AI Complementarity

Human-AI complementarity depends on employees' ability to recognize when AI should lead, when human expertise should dominate, and when both inputs should be integrated. This requires task allocation, role clarity, calibrated trust, and coordination (Ren et al., 2023; Dittmar & Sposato, 2026; Jain et al., 2023). Leadership is therefore consequential because leaders shape discretion, experimentation, knowledge sharing, performance expectations, and the extent to which AI is framed as a substitute, tool, or collaborative resource. Paradoxical leadership should facilitate complementarity because its underlying logic communicates that apparently contradictory expectations can coexist. Leaders can encourage employees to exploit AI's analytical and generative capabilities while continuing to contribute domain expertise and contextual judgment. They can also combine experimentation with standards and autonomy with accountability, thereby creating the structured flexibility needed for effective human-AI work. Emerging evidence supports this direction: paradoxical leadership shapes human-AI interaction and is associated with knowledge-related and technology-related behaviors that are useful for augmentation (Adeel et al., 2026; Batool, Ibrahim, & Adeel, 2024; Nagarajan, 2024; Zhang & Wu, 2026).

Proposition 1: *Paradoxical leadership positively facilitates human-AI complementarity by enabling employees to integrate AI capabilities with distinct human expertise, judgment, and contextual knowledge.*

3.2. Paradoxical Leadership and Responsible AI Agency

AI-enabled work can increase performance while simultaneously shifting cognitive authority away from employees. Excessive delegation may create automation bias, reduced verification, and declining independent judgment (Goddard et al., 2012; Klingbeil et al., 2024; Romeo & Conti, 2026). Paradoxical leadership offers a useful response because it does not equate empowerment with the removal of structure. Leaders can grant discretion to experiment with AI while maintaining expectations for validation, accountability, and final human responsibility. This both-and approach should preserve agency by communicating that AI use and human judgment are mutually necessary rather than mutually exclusive. Employees are encouraged to utilize AI actively while remaining answerable for the quality and appropriateness of resulting actions. This reasoning is consistent with responsible AI and human-in-the-loop research emphasizing meaningful human oversight and intervention (Papagiannidis et al., 2025; Lazaros et al., 2026; Sturm et al., 2026).

Proposition 2: *Paradoxical leadership positively facilitates responsible AI agency by encouraging employees to use AI while simultaneously retaining critical judgment, meaningful decision authority, and accountability for AI-assisted actions.*

3.3. Human-AI Complementarity and Innovative Work Behavior

Innovative work behavior involves the generation and implementation of useful new ideas (Scott & Bruce, 1994). Human-AI complementarity can enhance these activities by broadening the information available for ideation, accelerating search and comparison, and allowing employees to allocate more cognitive resources to interpretation and experimentation. Existing studies associate human-AI collaboration with creativity, innovation, and performance when capabilities are appropriately matched and coordinated (Kumar et al., 2025; Huang et al., 2026; Aslam, Liu, & Rahman, 2026). Complementarity should therefore improve innovative work behavior not because employees use more AI, but because distinct human and AI capabilities are combined in ways that improve both novelty and usefulness. This distinction is important because AI intensity without complementarity can lead to homogenization or dependence (Doshi & Hauser, 2024; Mascareño et al., 2026).

Proposition 3: *Human-AI complementarity is positively associated with employee innovative work behavior.*

3.4. Responsible AI Agency and Innovative Work Behavior

Innovation requires employees to evaluate, adapt, promote, and implement ideas rather than merely generate alternatives. Consequently, independent judgment remains essential even when GenAI expands the supply of potential solutions. AI can support decision quality when used reflectively, yet excessive delegation can produce automation bias, deskilling, and reduced reasoning (Albashrawi, 2025; Al-Anezi, 2026; Hao, Demir, & Eysers, 2024). Responsible AI agency should strengthen innovative work behavior because employees remain active contributors to AI-supported creative processes. They can distinguish useful novelty from generic or inappropriate suggestions, combine AI-generated information with contextual knowledge, and retain responsibility for implementation. In contrast, dependence can encourage cognitive inertia and conformity and thereby reduce creativity (Yin et al., 2024; Cui et al., 2026).

Proposition 4: *Responsible AI agency is positively associated with employee innovative work behavior.*

3.5. The Joint Role of Human-AI Complementarity and Responsible AI Agency

Paradox theory suggests that the greatest value should arise from the joint presence of capability integration and human agency. High complementarity accompanied by low agency may generate short-term efficiency but risks turning augmentation into substitution. Conversely, high agency accompanied by low complementarity may preserve human control while failing to realize AI's distinctive informational and computational advantages. The theoretically preferable configuration is therefore high complementarity combined with high responsible AI agency.

Under this condition, employees exploit AI's information-processing and generative capabilities while continuing to provide contextual judgment, critical evaluation, and accountability. This configuration represents the most complete both-and response to the human-AI paradox. The proposed interaction is a theoretical deduction from the combined paradox, complementarity, agency, and overreliance literatures rather than a relationship already established empirically.

Proposition 5: *Human-AI complementarity and responsible AI agency interact positively such that employee innovative work behavior is highest when both human-AI complementarity and responsible AI agency are high.*

3.6. AI Literacy as a Boundary Condition

Paradoxical leadership may encourage employees to both experiment with AI and preserve human judgment, but followers need sufficient AI competence to act on these expectations. Employees with limited AI literacy may struggle to identify appropriate task allocation, evaluate generated outputs, or recognize where human expertise adds distinctive value. In contrast, highly AI-literate employees should better understand AI's capabilities and limitations and translate both AI and leadership expectations into productive complementarity (Chee et al., 2025; Benlian & Pinski, 2025; Liu et al., 2025). Evidence linking AI literacy with performance and innovation outcomes reinforces this reasoning (Steinhauser & Heid, 2026; Zhou et al., 2026; Tan et al., 2026). AI literacy should therefore strengthen the ability of paradoxical leadership to generate genuine capability integration rather than superficial or inappropriate AI use.

Proposition 6: *AI literacy strengthens the positive relationship between paradoxical leadership and human-AI complementarity, such that the relationship is stronger at higher levels of AI literacy.*

3.7. AI Overreliance as a Boundary Condition

AI overreliance can undermine the innovation benefits of complementarity by shifting the relationship from augmentation toward substitution. Appropriate reliance may improve efficiency and decision performance, yet excessive deference can reduce critical evaluation, increase conformity, and weaken originality (Klingbeil et al., 2024; Romeo & Conti, 2026; Cui et al., 2026). Innovation is especially vulnerable because employees must challenge assumptions, discriminate among alternatives, and stay cognitively engaged in refining ideas. Accordingly, complementarity's positive effect on innovation should be weaker when employees defer excessively to AI. Under high overreliance, the apparent integration of human and AI inputs may increasingly reflect machine dominance rather than reciprocal complementarity.

Proposition 7: *AI overreliance weakens the positive relationship between human-AI complementarity and employee innovative work behavior, such that the relationship is weaker at higher levels of AI overreliance.*

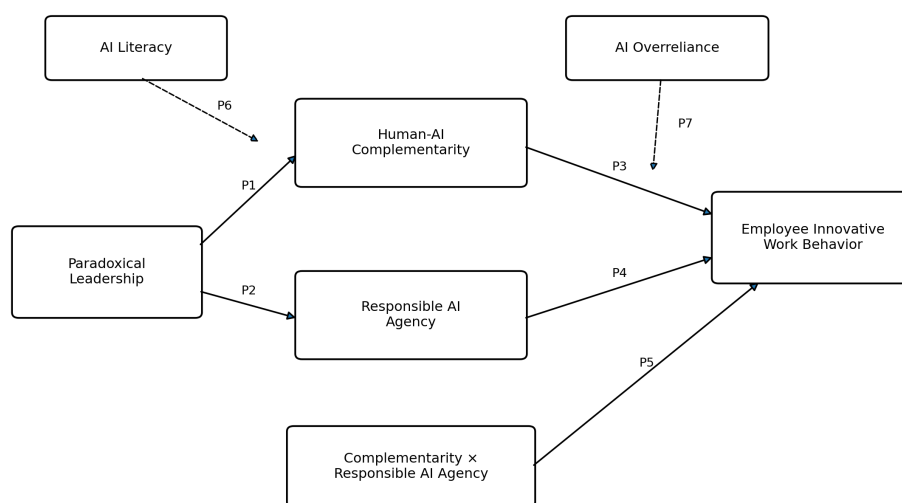


Figure 1: Conceptual framework of paradoxical leadership in human-AI collaboration.

4. Theoretical Contributions

The proposed framework contributes to paradox theory by conceptualizing the relationship between human agency and machine intelligence as an emerging organizational paradox. Traditional paradox research has largely addressed contradictions among organizational goals, roles, structures, identities, and human actors (Lewis, 2000; Smith & Lewis, 2011; Schad et al., 2016). GenAI extends this domain because one side of the tension can now involve an intelligent system capable of generating information, recommendations, and creative outputs. The framework therefore moves paradox theory beyond predominantly human-centered contradictions and toward the simultaneous management of human and machine intelligence. The framework also extends paradoxical leadership theory by proposing human-AI complementarity and responsible AI agency as mechanisms through which paradoxical leadership may influence innovation in AI-enabled work. Existing research establishes ambidexterity as an important pathway from paradoxical leadership to innovation (Zhang et al., 2022), while direct AI-related evidence remains emergent (Zhang & Wu, 2026). The present model therefore explains what paradoxical leadership may enable employees to do in hybrid work systems: integrate distinctive human and AI capabilities while preserving critical human judgment and accountability.

A third contribution is to distinguish human-AI complementarity from the mere presence, frequency, or intensity of AI use. Human-AI collaboration can enhance performance and creativity, but these benefits depend on task fit, role clarity, trust calibration, and coordination (Kumar et al., 2025; Aman et al., 2025; Dittmar & Sposato, 2026). More AI interaction does not necessarily imply stronger complementarity, particularly when dependence or homogenization emerges (Doshi & Hauser,

2024; Mascareño et al., 2026). The framework therefore shifts attention from the quantity of AI use to the quality of capability integration.

Fourth, the paper introduces responsible AI agency as a human-control mechanism that connects leadership, human agency, and responsible AI governance. Existing research addresses autonomy, oversight, critical evaluation, and accountability as related but often separate issues (Elmrayyan et al., 2026; Lazaros et al., 2026; Papagiannidis et al., 2025; Sturm et al., 2026). By integrating these ideas at the employee level, responsible AI agency explains whether humans remain meaningful, critical, and accountable actors within AI-enabled collaboration. This distinction complements human-AI complementarity: complementarity concerns whether capabilities work effectively together, whereas responsible AI agency concerns whether humans remain active evaluators and decision owners within that collaboration.

Finally, the framework advances an augmentation-based explanation of AI-enabled innovation. Rather than asking whether GenAI is inherently beneficial or harmful, it explains mixed findings by distinguishing augmentation from substitution. GenAI should contribute most strongly to innovation when it expands employees' cognitive resources while employees continue to provide contextual knowledge, critical evaluation, originality, and implementation judgment (Jia et al., 2024; Huang et al., 2026). The inclusion of AI literacy and AI overreliance further explains why technologically capable collaboration may generate innovation in some settings but dependence or diminished human contribution in others.

5. Future Research Agenda

The proposed framework provides several focused opportunities for empirical research. First, future studies should test whether human-AI complementarity and responsible AI agency operate as distinct mechanisms linking paradoxical leadership with employee innovative work behavior. Multi-wave and multisource designs would be particularly valuable because existing paradoxical leadership research demonstrates the usefulness of time-separated and multisource approaches for examining leadership mechanisms and innovation (Zhang et al., 2022). Because responsible AI agency is proposed here as a new conceptual construct, future research should also develop and validate its measurement while establishing discriminant validity from AI literacy, autonomy, trust, critical AI use, and AI dependence.

Second, future research should investigate the dynamic nature of paradoxical leadership in AI-enabled work. Zhang et al. (2022) explicitly call for greater attention to how paradoxical leaders adjust competing behaviors according to situational requirements. This issue is especially relevant for GenAI because the appropriate balance between human and machine contributions may change across work stages. Extensive AI involvement may support idea generation, whereas stronger human evaluation and accountability may be required during idea selection and implementation. Longitudinal, experimental, and experience-sampling designs could therefore examine when leaders should emphasize AI experimentation and when greater human oversight becomes necessary.

Finally, the proposed boundary conditions warrant further examination. AI literacy appears to enable employees to use and evaluate AI effectively, whereas excessive reliance can weaken critical judgment and increase conformity (Liu et al., 2025; Klingbeil et al., 2024; Romeo & Conti, 2026). Future research should determine whether different dimensions of AI literacy strengthen complementarity differently and whether AI overreliance produces linear or nonlinear effects. Extending the framework to team and organizational levels would also be valuable because leaders may shape shared human-AI norms, while organizational governance may determine whether employees retain meaningful oversight and accountability.

6. Practical Implications

The framework suggests that organizations should avoid treating GenAI implementation solely as a technology-adoption problem. Leaders must manage the tension between technological efficiency and human judgment. Paradoxical leadership is relevant because it supports a both-and approach in which employees can experiment with AI while maintaining autonomy, accountability, and clear performance expectations. Prior research similarly recommends developing paradoxical leadership capabilities to help managers combine flexibility with work requirements and autonomy with control (Zhang et al., 2022). Managers should also design work around human-AI complementarity rather than substitution. Employees should be encouraged to use AI where it provides distinctive advantages while retaining human responsibility for contextual interpretation, critical evaluation, and final judgment. Effective collaboration depends on task allocation, role clarity, trust calibration, and workflow design, not simply increasing AI use (Jain et al., 2023; Kumar et al., 2025). Organizations should therefore establish clear expectations concerning when AI outputs should be accepted, verified, modified, or rejected. Finally, organizations should invest in AI literacy while actively preventing AI overreliance. AI literacy enables employees to use AI effectively and critically and is increasingly associated with performance and innovation (Liu et al., 2025; Steinhäuser & Heid, 2026; Zhou et al., 2026). At the same time, excessive reliance can weaken judgment and encourage automation bias (Klingbeil et al., 2024; Romeo & Conti, 2026). AI training should therefore extend beyond technical proficiency to include verification, critical evaluation, responsible use, and clear human accountability.

7. Conclusion

The increasing integration of GenAI into organizational work creates a fundamental tension between exploiting machine intelligence and preserving human judgment, autonomy, creativity, and accountability. Drawing on paradox theory, this paper argues that these tensions should not be treated as problems to be resolved through an either-or choice. Instead, they require a both-and leadership approach that sustains the benefits of AI while protecting meaningful human contribution. The framework proposes that paradoxical leadership supports employee innovation through two complementary mechanisms: human-AI complementarity and responsible AI agency. The former captures the productive integration of distinct human and AI capabilities, while the latter emphasizes preserving critical judgment, decision authority, and accountability in AI-assisted work. By further considering AI literacy and AI overreliance as opposing boundary conditions, the framework explains why GenAI may enhance innovation in some settings while producing dependence or reduced human contribution in others. Overall, sustainable AI-enabled innovation depends not on maximizing AI use, but on leading human and machine intelligence as complementary yet tension-filled resources.

Acknowledgment Statement: The authors thank all participants and reviewers for their comments in helping complete this manuscript.

Conflicts of interest: The authors declare that they have no known competing financial interests or personal relationships that could have appeared to influence the work reported in this paper.

Funding statements: No external funding was received for this research; the study was conducted without financial support from any funding agency or organisation.

Data availability statement: This is a conceptual study. Sources are available through the reference list.

Disclaimer: The views and opinions expressed in this article are those of the author(s) and contributor(s) and do not necessarily reflect IMF's or editors' official policy or position. All liability for harm done to individuals or property as a result of any ideas, methods, instructions, or products mentioned in the content is expressly disclaimed.

Declaration of generative AI and AI-assisted technologies in the writing process: During the preparation of this work, the author(s) used ChatGPT for copy-editing. After using this tool/service, the author(s) reviewed and edited the content as needed and take(s) full responsibility for the content of the publication.

References

- Adeel, A., Ashraf, A. A., & Sheikh, A. A. (2026). Why AI does not always enhance employee creativity: Metacognitive calibration and AI engagement in human–AI innovation systems. *International Journal of Systematic Innovation*, 10(4), 026120027. [https://doi.org/10.6977/IJoSI.202608_10\(4\).0011](https://doi.org/10.6977/IJoSI.202608_10(4).0011)
- Adeel, A., Hussain, G., & Sheikh, A. A. (2026). How Culture Is Communicatively Constituted In Intercultural Healthcare Organizations: Leadership Sensegiving, Knowledge Sharing, And Employee Innovation. *Journal of Intercultural Communication*, 26(1), 32–43. <https://doi.org/10.36923/jicc.v26i1.1409>
- Ahmad, I., & Mehmood, S. (2025). Rituals of Openness: Vulnerability Practices in Multidisciplinary Professional Settings Beyond Healthcare. *Innovation Journal of Social Sciences and Economic Review*, 7(1), 50–63. <https://doi.org/10.36923/ijsser.v7i1.297>
- Ahmad, I., Thurasamy, R., Adeel, A., & Alam, B. (2023). Promotive Voice, Leader-member Exchange, and Creativity Endorsement: The Role of Supervisor-Attributed Motives. *Journal of Intercultural Communication*, 23(3), 1–13. <https://doi.org/10.36923/jicc.v23i3.121>
- Al-Anezi, F. M. (2026). Generative artificial intelligence in healthcare: Automation bias, deskilling, and cognitive implications: A systematic review. *Journal of Healthcare Leadership*. Advance online publication. <https://doi.org/10.2147/JHL.S590498>
- Alamanda, D. T., Ahmed, A., Kurniady, D. A., Rahayu, A., Ahmad, I., & Hashim, N. A. A. N. (2024). Linking Gender To Creativity: Role of Risk Taking and Support For Creativity Towards Creative Potential of Employees. *Journal of Intercultural Communication*, 24(1), 1–17. <https://doi.org/10.36923/jicc.v24i1.219>
- Albashrawi, M. (2025). Generative AI for decision-making: A multidisciplinary perspective. *Journal of Innovation and Knowledge*. Advance online publication. <https://doi.org/10.1016/j.jik.2025.100751>
- Alon-Barkat, S., & Busuioc, M. (2023). Human-AI interactions in public sector decision making: Automation bias and selective adherence to algorithmic advice. *Journal of Public Administration Research and Theory*, 33(3), 449–466. <https://doi.org/10.1093/jopart/muac007>
- Alshalfan, A. (2026). Change management for AI in service operations: A paradox and routine dynamics perspective. *Journal of Global Information Management*. Advance online publication. <https://doi.org/10.4018/JGIM.416702>
- Aman, F., Hamid, S. R. A., Mat-Isa, A. M., & Mohamad, M. H. S. (2025). A systematic review of trends in human-AI collaboration research. *International Journal of Technology, Knowledge and Society*, 21(1), 189–217. <https://doi.org/10.18848/1832-3669/CGP/v21i01/189-217>
- Annappureddy, R., Fornaroli, A., & Gatica-Pérez, D. (2025). Generative AI literacy: Twelve defining competencies. *Digital Government: Research and Practice*. Advance online publication. <https://doi.org/10.1145/3685680>
- Aslam, Y., Liu, Z., & Rahman, M. M. (2026). The synergy of artificial intelligence capabilities and creativity: Exploring positive outcomes through qualitative research. *SAGE Open*. Advance online publication. <https://doi.org/10.1177/21582440261422023>
- Batool, S., Ibrahim, H. I., & Adeel, A. (2024). How responsible leadership pays off: Role of organizational identification and organizational culture for creative idea sharing. *Sustainable Technology and Entrepreneurship*, 3(2), 100057. <https://doi.org/10.1016/j.stae.2023.100057>
- Batool, U., Raziq, M. M., & Sarwar, N. (2023). The paradox of paradoxical leadership: A multi-level conceptualization. *Human Resource Management Review*, 33(4), 100983. <https://doi.org/10.1016/j.hrmr.2023.100983>
- Benlian, A., & Pinski, M. (2025). The AI literacy development canvas: Assessing and building AI literacy in organizations. *Business Horizons*. Advance online publication. <https://doi.org/10.1016/j.bushor.2025.10.001>
- Boemelburg, R., Zimmermann, A., & Palmié, M. (2023). How paradoxical leaders guide their followers to embrace paradox: Cognitive and behavioral mechanisms of paradox mindset development. *Long Range Planning*, 56(5), 102319. <https://doi.org/10.1016/j.lrp.2023.102319>
- Chee, H., Ahn, S., & Lee, J. (2025). A competency framework for AI literacy: Variations by different learner groups and an implied learning pathway. *British Journal of Educational Technology*. Advance online publication. <https://doi.org/10.1111/bjet.13556>
- Cheng, Q., Dai, Y., Liu, X., & Peng, S. (2026). The trust crisis in artificial intelligence: AI hallucinations and human-AI collaboration. *Technology in Society*, 86, 103286. <https://doi.org/10.1016/j.techsoc.2026.103286>
- Choi, Y. T., Park, C. W., Paydas Turan, C. P., & Lee, H. (2026). Beyond the creativity paradox: A theory-informed framework for role-based integration of generative AI in organisational creativity. *Information Systems Frontiers*. Advance online publication. <https://doi.org/10.1007/s10796-026-10746-y>

- Cristofaro, M., & Bañón-Gómis, A. J. (2026). Dancing with the algorithm: A framework to navigate knowledge and autonomy in AI-assisted managerial decisions. *Journal of Knowledge Management*. Advance online publication. <https://doi.org/10.1108/JKM-06-2025-0870>
- Cui, S., Wang, L., Cao, W., & Zhu, T. (2026). Gain or loss? The dual effects of dependence on AI on employee's creativity. *International Journal of Information Management*, 86, 103001. <https://doi.org/10.1016/j.ijinfomgt.2025.103001>
- Di Plinio, S. (2025). Panta Rh-AI: Assessing multifaceted AI threats on human agency and identity. *Social Sciences and Humanities Open*, 13, 101434. <https://doi.org/10.1016/j.ssaho.2025.101434>
- Dittmar, E. C., & Sposato, M. (2026). Optimising human-AI decision performance: A trust and capability framework for knowledge management. *Knowledge and Process Management*. Advance online publication. <https://doi.org/10.1002/kpm.70059>
- Doshi, A. R., & Hauser, O. P. (2024). Generative AI enhances individual creativity but reduces the collective diversity of novel content. *Science Advances*, 10(28), eadn5290. <https://doi.org/10.1126/sciadv.adn5290>
- Elmrayyan, N., Bani-Hani, I., & Alawadi, S. (2026). Generative AI and decision autonomy: A framework of psychological and structural empowerment. *Journal of Decision Systems*. Advance online publication. <https://doi.org/10.1080/12460125.2026.2657495>
- Esposito De Falco, S., Laviola, F., Mercuri, F., & Cucari, N. (2026). Different routes, same storm: A three-dimensional paradox view of generative AI's governance. *Management Decision*. Advance online publication. <https://doi.org/10.1108/MD-10-2025-3328>
- Fügener, A., Walzner, D. D., & Gupta, A. (2025). Roles of artificial intelligence in collaboration with humans: Automation, augmentation, and the future of work. *Management Science*. Advance online publication. <https://doi.org/10.1287/mnsc.2024.05684>
- García-Barrios, D. A. (2025). Epistemic partner or cognitive crutch: A conceptual model of cognitive offloading in human-artificial intelligence collaboration. *IEEE Internet Computing*. Advance online publication. <https://doi.org/10.1109/MIC.2025.3626694>
- Goddard, K., Roudsari, A., & Wyatt, J. C. (2012). Automation bias: A systematic review of frequency, effect mediators, and mitigators. *Journal of the American Medical Informatics Association*, 19(1), 121-127. <https://doi.org/10.1136/amiajnl-2011-000089>
- Hao, X., Demir, E., & Eysers, D. (2024). Exploring collaborative decision-making: A quasi-experimental study of human and generative AI interaction. *Technology in Society*, 79, 102662. <https://doi.org/10.1016/j.techsoc.2024.102662>
- Hillebrand, L., Raisch, S., & Schad, J. (2025). Managing with artificial intelligence: An integrative framework. *Academy of Management Annals*. Advance online publication. <https://doi.org/10.5465/annals.2022.0072>
- Huang, S., Long, L., Zhu, Y., & Zhu, J. N. Y. (2026). Human-GenAI collaboration across creative phases: Cognitive mechanisms shaping novelty and usefulness. *International Journal of Information Management*, 86, 102986. <https://doi.org/10.1016/j.ijinfomgt.2025.102986>
- Jain, R., Garg, N., & Khera, S. N. (2023). Effective human-AI work design for collaborative decision-making. *Kybernetes*. Advance online publication. <https://doi.org/10.1108/K-04-2022-0548>
- Jia, N., Luo, X., Fang, Z., & Liao, C. (2024). When and how artificial intelligence augments employee creativity. *Academy of Management Journal*, 67(1), 5-33. <https://doi.org/10.5465/amj.2022.0426>
- Klingbeil, A., Grütznier, C., & Schreck, P. (2024). Trust and reliance on AI: An experimental study on the extent and costs of overreliance on AI. *Computers in Human Behavior*, 160, 108352. <https://doi.org/10.1016/j.chb.2024.108352>
- Kumar, N., Kumar, R. R., & Raj, A. (2025). Establishing antecedents and outcomes of human-AI collaboration: Meta-analysis. *Journal of Computer Information Systems*. Advance online publication. <https://doi.org/10.1080/08874417.2025.2492885>
- Lazaros, K., Vrahatis, A. G., & Kotsiantis, S. (2026). Human-in-the-loop artificial intelligence: A systematic review of concepts, methods, and applications. *Entropy*, 28(4), 377. <https://doi.org/10.3390/e28040377>
- Lee, B. C., & Chung, J. (2024). An empirical investigation of the impact of ChatGPT on creativity. *Nature Human Behaviour*, 8(8), 1540-1552. <https://doi.org/10.1038/s41562-024-01953-1>
- Lee, H.-K., Yoon, C., & Lee, B. G. (2026). Operationalising human-centred AI governance under the EU AI Act: A governance framework for human oversight and data accountability. *Systems*, 14(7), 849. <https://doi.org/10.3390/systems14070849>
- Lewis, M. W. (2000). Exploring paradox: Toward a more comprehensive guide. *Academy of Management Review*, 25(4), 760-776. <https://doi.org/10.2307/259204>
- Liu, X., Zhang, L., & Wei, X. (2025). Generative artificial intelligence literacy: Scale development and its effect on job performance. *Behavioral Sciences*, 15(6), 811. <https://doi.org/10.3390/bs15060811>
- Liu, Y., Xu, S., & Zhang, B. (2020). Thriving at work: How a paradox mindset influences innovative work behavior. *Journal of Applied Behavioral Science*, 56(4), 454-481. <https://doi.org/10.1177/0021886319888267>
- Ma, L., Ge, L., Li, Y., & Wang, Y. (2026). When employees meet GenAI: The double-edged effects of AI attitude on creativity. *Technovation*, 142, 103495. <https://doi.org/10.1016/j.technovation.2026.103495>
- Markus, A., Carolus, A., & Wienrich, C. (2025). Objective measurement of AI literacy: Development and validation of the AI competency objective scale (AICOS). *Computers and Education: Artificial Intelligence*, 8, 100485. <https://doi.org/10.1016/j.caeai.2025.100485>
- Marvi, R., Foroudi, P., & AmirDadbar, N. (2025). Dynamics of user engagement: AI mastery goal and the paradox mindset in AI-employee collaboration. *International Journal of Information Management*, 85, 102908. <https://doi.org/10.1016/j.ijinfomgt.2025.102908>
- Mascareño, J., Wörtler, B., Przegalińska, A., & Ciechanowski, L. (2026). When proximal collaboration with AI hinders innovation: The moderating role of idea originality and reliance on AI. *Computers in Human Behavior Reports*, 13, 101054. <https://doi.org/10.1016/j.chbr.2026.101054>
- Nagarajan, N. C. (2024). Paradoxical leadership and employee creativity: Knowledge sharing and hiding as mediators. *Journal of Knowledge Management*, 28(5), 1391-1415. <https://doi.org/10.1108/JKM-10-2022-0779>

- Novelli, C., Taddeo, M., & Floridi, L. (2024). Accountability in artificial intelligence: What it is and how it works. *AI & Society*, 39(1), 1-14. <https://doi.org/10.1007/s00146-023-01635-y>
- Okamura, K., & Yamada, S. (2020). Adaptive trust calibration for human-AI collaboration. *PLOS ONE*, 15(2), e0229132. <https://doi.org/10.1371/journal.pone.0229132>
- Özer, M., Perc, M., & Özçelik, H. T. (2026). Human agency and epistemic authority under generative artificial intelligence. *Open Praxis*, 18(3), 1149. <https://doi.org/10.55982/openpraxis.18.3.1149>
- Papadonikolaki, E. (2026). Rethinking artificial intelligence in projects: A paradox lens. *International Journal of Project Management*, 44(4), 102857. <https://doi.org/10.1016/j.ijproman.2026.102857>
- Papagiannidis, E., Mikalef, P., & Conboy, K. (2025). Responsible artificial intelligence governance: A review and research framework. *Journal of Strategic Information Systems*, 34(1), 101885. <https://doi.org/10.1016/j.jsis.2024.101885>
- Parasuraman, R., & Manzey, D. H. (2010). Complacency and bias in human use of automation: An attentional integration. *Human Factors*, 52(3), 381-410. <https://doi.org/10.1177/0018720810376055>
- Raisch, S., & Krakowski, S. (2021). Artificial intelligence and management: The automation-augmentation paradox. *Academy of Management Review*, 46(1), 192-210. <https://doi.org/10.5465/amr.2018.0072>
- Ren, Y., Deng, X. N., & Joshi, K. D. (2023). Unpacking human and AI complementarity: Insights from recent works. *Data Base for Advances in Information Systems*, 54(3), 39-52. <https://doi.org/10.1145/3614178.3614180>
- Ritala, P., Ruokonen, M., & Ramaul, L. (2024). Transforming boundaries: How does ChatGPT change knowledge work? *Journal of Business Strategy*, 45(4), 251-260. <https://doi.org/10.1108/JBS-05-2023-0094>
- Romeo, G., & Conti, D. (2026). Exploring automation bias in human-AI collaboration: A review and implications for explainable AI. *AI & Society*. Advance online publication. <https://doi.org/10.1007/s00146-025-02422-7>
- Schad, J., Lewis, M. W., Raisch, S., & Smith, W. K. (2016). Paradox research in management science: Looking back to move forward. *Academy of Management Annals*, 10(1), 5-64. <https://doi.org/10.5465/19416520.2016.1162422>
- Scott, S. G., & Bruce, R. A. (1994). Determinants of innovative behavior: A path model of individual innovation in the workplace. *Academy of Management Journal*, 37(3), 580-607. <https://doi.org/10.2307/256701>
- Smith, W. K., & Lewis, M. W. (2011). Toward a theory of paradox: A dynamic equilibrium model of organizing. *Academy of Management Review*, 36(2), 381-403. <https://doi.org/10.5465/AMR.2011.59330958>
- Steinhauser, S., & Heid, P. (2026). Organizational readiness, AI literacy, and the new frontier of R&D: How generative AI shapes innovation capacity. *R&D Management*. Advance online publication. <https://doi.org/10.1111/radm.70038>
- Sturm, S., Krause, F., & van Giffen, B. (2026). Beyond control? Governing AI in organizations for responsible use. *European Management Journal*. Advance online publication. <https://doi.org/10.1016/j.emj.2026.06.003>
- Tan, W., Nawaz, M., Shu, T., & Ramzan, B. (2026). Collaborative artificial intelligence literacy and employee performance: Task-technology fit and technostress in a moderated mediation model. *South African Journal of Business Management*, 57(1), 5903. <https://doi.org/10.4102/sajbm.v57i1.5903>
- Vössing, M., Kuhl, N., Lind, M., & Satzger, G. (2022). Designing transparency for effective human-AI collaboration. *Information Systems Frontiers*, 24(3), 877-895. <https://doi.org/10.1007/s10796-022-10284-3>
- Wei, W., Zhou, Y., & Wang, D. (2023). Learning to integrate conflicts: Paradoxical leadership fosters team innovation. *Journal of Business Research*, 165, 114076. <https://doi.org/10.1016/j.jbusres.2023.114076>
- Yin, M., Jiang, S., & Niu, X. (2024). Can AI really help? The double-edged sword effect of AI assistant on employees' innovation behavior. *Computers in Human Behavior*, 150, 107987. <https://doi.org/10.1016/j.chb.2023.107987>
- Zhang, H., Zhu, L., Zhang, A., & Shohruh, K. (2026). The influence of generative artificial intelligence usage on employees' innovative job performance. *PLOS ONE*, 21(3), e0327786. <https://doi.org/10.1371/journal.pone.0327786>
- Zhang, M. J., Zhang, Y., & Law, K. S. (2022). Paradoxical leadership and innovation in work teams: The multilevel mediating role of ambidexterity and leader vision as a boundary condition. *Academy of Management Journal*, 65(5), 1652-1679. <https://doi.org/10.5465/amj.2017.1265>
- Zhang, Y., Waldman, D. A., Han, Y. L., & Li, X. B. (2015). Paradoxical leader behaviors in people management: Antecedents and consequences. *Academy of Management Journal*, 58(2), 538-566. <https://doi.org/10.5465/amj.2012.0995>
- Zhang, Z., & Wu, Y. (2026). Different types of human-AI interaction and employees' knowledge territorial behavior: The role of organizational dehumanization and paradoxical leadership. *Frontiers in Psychology*, 17, 1854239. <https://doi.org/10.3389/fpsyg.2026.1854239>
- Zhou, J., Lü, Y., & Li, W. (2026). Human-centric sustainability in organizations: The role of GAI literacy in individual well-being and innovative performance. *Telecommunications Policy*, 50(4), 103273. <https://doi.org/10.1016/j.telpol.2026.103273>

About the Author(s).

Ahmad Adeel is an associate professor of research at the Faculty of Business and Communications, INTI International University, Persiaran Perdana BBN, Putra Nilai 71800 Nilai, Negeri Sembilan, Malaysia. He received his PhD in Business Administration from Huazhong University of Science and Technology, Wuhan, China. His research interests are conflict management, motivation, and employee creativity.

Samreen Batool is a PhD candidate at the School of Management, Universiti Sains Malaysia (USM), Minden, 11800 Penang, Malaysia.